

A S S E M B L É E N A T I O N A L E

X V ^e L É G I S L A T U R E

Compte rendu

Commission des affaires sociales

- Réunion, conjointe avec la commission des lois, sur le thème de l'intelligence artificielle en présence de M. Alain Bensoussan, avocat, président de l'association du droit des robots, M. Grégoire Loiseau, Professeur de droit privé à l'université Paris I-La Sorbonne, M. Jacques Lucas, médecin, vice-président du conseil national de l'ordre des médecins..... 2
- Présences en réunion 18

Mercredi

6 juin 2018

Séance de 11 heures

Compte rendu n° 79

SESSION ORDINAIRE DE 2017-2018

**Présidence de
Mme Brigitte Bourguignon,
Présidente,
et de
Mme Yaël Braun-Pivet,
Présidente de la
commission des lois**



COMMISSION DES AFFAIRES SOCIALES

Mercredi 6 juin 2018

La séance est ouverte à onze heures.

*(Présidence de Mme Brigitte Bourguignon, présidente de la Commission
et de Mme Yaël Braun-Pivet, présidente de la commission des lois)*

La commission des affaires sociales organise, conjointement avec la commission des lois, une réunion sur le thème de l'intelligence artificielle en présence de M. Alain Bensoussan, avocat, président de l'association du droit des robots, M. Grégoire Loiseau, professeur de droit privé à l'université Paris I-La Sorbonne, M. Jacques Lucas, médecin, vice-président du conseil national de l'ordre des médecins.

Mme la présidente Yaël Braun-Pivet. Cette table ronde a pour thématique l'« intelligence artificielle », c'est-à-dire l'ensemble des technologies ayant pour objectif de comprendre comment fonctionne la cognition humaine et comment la reproduire, voire la dépasser. L'intelligence artificielle trouve son origine au milieu du XX^e siècle, mais ce n'est qu'au début des années 2010, sous l'impulsion de l'explosion des puissances de calcul, de la multiplication des volumes de données et de l'essor de nouveaux algorithmes, qu'elle a pris sa nouvelle dimension : celle d'une révolution technologique, mais aussi économique et sociétale. Le développement de l'intelligence artificielle représente en effet un bouleversement de nature à transformer profondément la société : les applications actuelles et futures touchent à tous les domaines de la vie quotidienne, avec des conséquences potentielles considérables. Je pense par exemple à la santé, qui connaît un recours croissant aux robots, en particulier en neurochirurgie, aux transports, avec la voiture autonome, ou encore aux loisirs, avec les assistants virtuels.

Le développement de l'intelligence artificielle est soumis à une contrainte d'acceptabilité sociale forte, qui est liée aux représentations, parfois catastrophiques, dont elle fait l'objet, et à la crainte de possibles dérives, comme celles qui pourraient résulter de certains systèmes algorithmiques dont on comprend encore mal le fonctionnement. Porteur d'espoirs mais aussi de craintes, le développement de l'intelligence artificielle se caractérise aujourd'hui par la place prépondérante qu'occupe la recherche privée, laquelle est dominée par des entreprises américaines et chinoises.

Nous accueillons M. Jacques Lucas, vice-président du Conseil national de l'ordre des médecins, M. Alain Bensoussan, avocat, et M. Grégoire Loiseau, professeur de droit privé à l'université Paris 1 Panthéon-Sorbonne. Je vous invite, messieurs, à nous présenter à titre liminaire le cadre scientifique et juridique dans lequel s'inscrit l'intelligence artificielle et à nous indiquer quelles sont, à vos yeux, les pistes d'évolution possibles. Nous en viendrons ensuite à des questions, sous la forme d'interventions de deux minutes, à caractère interrogatif.

M. Jacques Lucas, médecin, vice-président du Conseil national de l'Ordre des médecins. Merci beaucoup, madame la présidente. Permettez-moi d'attirer l'attention de vos

deux commissions sur l'ouvrage intitulé « *Le médecin et le patient dans le monde des data, des algorithmes et de l'intelligence artificielle* », qui a été publié par le Conseil national de l'Ordre.

Tout va reposer, en fin de compte, sur la collecte et le traitement des données. Si elles ne sont pas de bonne qualité, l'intelligence artificielle sera sotte : elle sera nourrie de données non pertinentes. On estime aujourd'hui qu'environ 30 % des publications médicales et scientifiques aboutissent à des conclusions erronées parce que le traitement des données n'a pas été de bonne qualité. Avant même de parler d'intelligence artificielle, il faut donc regarder, ainsi que vous l'avez dit, madame la présidente, comment elle est constituée et si elle est véritablement intelligente. Ainsi que l'a souligné Jean-Gabriel Ganascia dans sa contribution à notre ouvrage, il importe de ne pas reproduire des certitudes dogmatiques lorsque l'on construit un algorithme : le traitement des données ne serait pas pertinent.

Par ailleurs, et je pense que l'Assemblée nationale ne peut qu'être sensible à cette question, il ne faut pas que les avancées scientifiques accentuent les discriminations en raison d'un inégal accès, notamment sur le plan territorial, aux dispositifs qui pourraient être mis à la disposition des citoyens et des professionnels de santé. J'utilise le terme de citoyen plutôt que celui de patient, car les dispositifs d'intelligence artificielle peuvent venir au secours de personnes qui ne sont pas malades, notamment dans le cadre du dépistage et de la prévention, ce qui nécessite bien évidemment une attention particulière.

Le Conseil national de l'ordre a produit 33 recommandations, d'inégale portée. L'une d'entre elles souligne que les technologies, et notamment celles de l'intelligence artificielle, doivent être au service de la personne et de la société. Celles-ci doivent être libres et non asservies aux géants technologiques, qui ne sont d'ailleurs pas des États mais des sociétés de droit privé. Ce principe éthique, qui est fondamental à nos yeux, doit être réaffirmé à l'heure où les dystopies et les utopies les plus excessives, notamment au sujet de l'homme immortel, sont largement médiatisées. C'est pourquoi nous recommandons l'adoption de règles protectrices dans le droit positif. Nous avons bien conscience que ces règles doivent avoir une portée internationale : la France et l'Europe politique doivent en faire une de leurs ambitions majeures. Les technologies sont au service d'un projet de société, et l'on doit réaffirmer les repères de ce qui fait notre humanité.

Guy Vallancien, qui a lui aussi contribué à notre livre blanc et que vous avez peut-être auditionné, considère en outre que si l'on est attentif au climat et au réchauffement climatique, il faut vraisemblablement organiser une conférence internationale sur le déploiement de l'intelligence artificielle.

Voilà les observations que je voulais formuler à titre liminaire. Je répondrai volontiers à des questions plus précises et plus concrètes si vous le souhaitez.

M. Grégoire Loiseau, professeur de droit privé à l'université Paris 1 Panthéon-Sorbonne. En ce qui concerne l'appréhension de l'intelligence artificielle par le droit, je souhaite attirer votre attention sur deux points.

Tout d'abord, rien n'existe pour l'instant, car il s'agit d'un phénomène émergent. Les questions sont surtout posées sous un angle économique ou du point de vue du droit social, qui n'est pas l'objet de cette table ronde. Les évolutions qui sont attendues dans le domaine de l'intelligence artificielle doivent nous conduire à adopter un certain nombre de mesures relevant de l'éthique et, peut-être, de la loi de bioéthique.

J'en viens maintenant aux risques de dérive. À titre personnel – je sais que ce point de vue n'est pas partagé par tous –, j'en vois deux.

Le premier risque de dérive est la personnification des « robots » – c'est le terme usuel, même s'il est souvent inapproprié, car l'intelligence artificielle peut prendre une forme immatérielle. On assiste déjà à une telle personnification pour certaines utilisations : les robots d'assistance à la personne et ceux de compagnie, qui suscitent une relation émotionnelle. Le pas suivant est la personnification juridique : on le voit bien avec le débat qui concerne les animaux. La personnification de certaines machines intelligentes, les plus sophistiquées, risque de donner lieu à la revendication d'une personnalité juridique, que l'on présente comme fonctionnelle ou technique, à l'image de celle des personnes morales. Or on sait très bien, en tout cas chez les juristes, que la reconnaissance de la personnalité juridique des personnes morales leur a permis d'accéder aux mêmes droits que les êtres humains, sur le terrain de l'égalité – Ripert le soulignait déjà au milieu du XX^e siècle. La reconnaissance d'une personnalité juridique, même technique, aboutirait à la revendication de droits pour certaines utilisations de l'intelligence artificielle.

Vous savez sans doute qu'une résolution adoptée l'an dernier par le Parlement européen demande notamment la reconnaissance, à terme, d'une personnalité électronique pour les robots. C'est un sujet qu'il faut prendre au sérieux. Un certain nombre de chercheurs – 150, me semble-t-il – ont adressé le mois dernier une lettre à la Commission européenne afin de l'alerter.

On ne peut pas traiter cette question en droit national, en faisant abstraction des autres droits : on ne peut agir utilement qu'au niveau de l'Union européenne. Néanmoins, la France a un rôle à jouer, ne serait-ce que parce qu'elle a été un des premiers pays européens à reconnaître la primauté de la personne humaine en adoptant la loi du 29 juillet 1994 relative à la bioéthique, qui n'a absolument pas vieilli, même s'il faut l'adapter, ce qui est bien différent. Il me semble que cette sensibilité historique à l'humain donne à la France une voix au chapitre sur ces questions.

Un deuxième risque de dérive est lié au développement d'un discours utopiste, qui est soutenu par des firmes, notamment américaines, sur l'utilisation des biotechnologies, des nanotechnologies et de l'intelligence artificielle pour augmenter la résilience de l'être humain, en créant un être humain « augmenté » sur le plan de ses capacités cognitives et physiques, jusqu'à l'immortalité. Même si on ne va pas jusque-là, l'idée est de perfectionner l'être humain. À titre individuel, qui serait contre la possibilité de vivre plus longtemps sans être affecté par les tourments de l'âge, comme la maladie d'Alzheimer ? Il faut néanmoins raisonner à l'échelle collective : c'est la conception même de l'être humain, de l'espèce humaine, qui est en jeu. Selon les tenants du courant le plus radical du transhumanisme, il faudrait transmettre génétiquement l'amélioration de l'être humain, ce que l'on est d'ailleurs en passe de faire.

Le droit européen et le droit français doivent se positionner. Depuis 1994, l'article 16-4 du code civil pose le principe de l'intégrité de l'espèce humaine. Je suggère, très modestement, d'y ajouter un principe d'intangibilité. Cela signifie que notre espèce doit connaître une évolution naturelle, et non pas une évolution artificielle reposant sur les nanobiotechnologies et l'intelligence artificielle, créature qui pourrait très facilement échapper à son créateur.

Je ne cherche pas à tenir des propos inquiétants, mais à faire prendre conscience des enjeux, à un moment où rien n'est encore fait, et où tout reste à faire. Le risque, me semble-t-il, est de se laisser dépasser par les évolutions en se disant qu'elles sont pour plus tard et que l'on va donc se limiter, dans un premier temps, aux apports en matière médicale et dans le domaine du transport, ou bien aux risques de suppression d'emplois : les aspects relevant du droit des personnes me paraissent aussi extrêmement importants.

M. Alain Bensoussan, avocat, président de l'Association du droit des robots.
Mon intervention concerne le cadre juridique de l'intelligence artificielle dans le domaine de la santé. Il me paraît important de souligner d'emblée qu'il y a une spécificité en la matière. On pourrait donc se placer soit dans un cadre général, soit dans celui d'une réglementation particulière.

De quoi parle-t-on au juste ? Un robot physique est une coque dotée de moyens mécatroniques et de capteurs. La révolution actuelle concerne davantage les capteurs que les technologies de l'intelligence artificielle, qui remontent aux années 1970. Un robot logiciel est un précipité d'algorithmes, lesquels sont des suites d'instructions finies qui permettent d'atteindre un résultat sans avoir à le démontrer. Le superbe rapport de Cédric Villani, que je salue, comporte une typologie des algorithmes qui donne une vision très précise de ce que l'on appelle « l'intelligence faible ».

Je vous propose de réfléchir à un cadre juridique pour cette intelligence artificielle faible, qui est celle d'aujourd'hui, mais aussi d'essayer d'envisager la question de l'intelligence artificielle moyenne. L'intelligence artificielle dite faible est à l'heure actuelle supérieure à l'humain pour la plupart des opérations, et elle comporte des applications spécifiques dans le domaine de la santé, notamment en neurologie et en dermatologie. Pour une finalité donnée, qui peut aller du jeu de go à l'urologie, les robots physiques comportant un robot logiciel intégré sont supérieurs à l'humain.

J'aborderai successivement trois thèmes : les algorithmes et la santé ; les robots et la santé ; les données et la santé.

En ce qui concerne le premier point, je pense qu'il y a trois règles à prendre en compte sur le plan juridique. En tant que praticien de l'intelligence artificielle, qui s'inscrit dans le droit « mou », je les applique déjà.

La première règle est le droit, pour l'individu, de savoir s'il est en face d'un robot ou d'un humain. Ce droit à la connaissance est un élément déterminant en matière de transparence. L'intelligence artificielle est en train de se disséminer dans l'ensemble des hôpitaux et des cliniques, mais aussi des usines et des entreprises, ainsi qu'à l'intérieur des domiciles.

Un autre élément déterminant est le droit à l'erreur. Le système actuel de responsabilité n'est pas applicable. Dans un cas de leucémie, Watson, qui est un système développé par IBM, a proposé au Japon une solution que le corps médical ne parvenait pas à trouver. Un robot logiciel est supérieur à l'humain pour certains cancers. Laurent Alexandre a ainsi montré que le système se trompe une fois sur dix, contre une fois sur deux pour les médecins. Si l'erreur robotique est inférieure à l'erreur humaine, le droit conçu pour cette dernière ne peut pas s'appliquer : on ne doit plus se trouver dans un régime de responsabilité pour faute, mais sans faute. Il faut abandonner la distinction entre l'obligation de moyens et l'obligation de résultat, car cela n'a guère de sens : il doit s'agir d'une responsabilité pour

dommages subis. Cela revient à généraliser la loi Badinter à l'ensemble de l'intelligence artificielle.

Le troisième droit fondamental est le droit à la certification. Les algorithmes dont nous parlons sont auto-apprenants. Des systèmes installés dans des cliniques ou des hôpitaux différents vont muter et ils n'auront donc pas le même comportement face à de nouveaux patients. La certification obligatoire permettrait d'avoir une présomption d'irresponsabilité. Comme en matière de presse, il faut un système en cascade – le robot, puis le concepteur de la plateforme d'intelligence artificielle et ensuite l'ensemble des éléments habituels – le fournisseur, le vendeur et tout ce que prévoit la directive « machines », laquelle est à mon avis totalement inapplicable. Dans le cadre d'un système de certification, la plateforme d'intelligence artificielle verrait sa responsabilité dérogée en termes d'approche fautive.

J'en viens à la question des robots et de la santé. J'ai lancé l'idée, dans le cadre du débat mondial qui s'est engagé, que les robots ont un droit à la souveraineté et qu'ils doivent être reconnus juridiquement. J'ai eu l'occasion de faire une présentation plus détaillée devant le Parlement, et je pense avoir convaincu différentes personnalités dans le monde... Si tous les humains sont des personnes, toutes les personnes ne sont pas humaines. Le concept de personne a évolué, notamment en ce qui concerne les personnes morales, mais il y a aussi une reconnaissance de la personne « fleuve » dans le domaine de la climatologie, et la personne « montagne » a aussi été reconnue. La personne, au sens de la personnalité juridique, n'est qu'un système : ce n'est qu'une marmite, un vecteur, un mot-valise. Pourquoi a-t-on affaire ici aussi à des personnes ? Parce que les systèmes sont apprenants, autonomes et mutants. Les personnes humaines sont libres, et l'on a créé des droits et des obligations pour elles, dans le cadre d'une personnalité juridique générale. Les personnes morales sont libres aussi, et l'on a créé une personnalité juridique particulière, au moyen d'une adaptation. Osons reconnaître que les personnes robots ont une personnalité juridique singulière en matière de responsabilité, de dignité, de traçabilité et de décision en dernier ressort.

Je vous propose trois droits. Le premier est un droit en matière d'intimité numérique, qui n'existe pas, à l'heure actuelle, au-delà du droit à la vie privée. Je crois que le rapport entre le patient et le robot est de même nature que celui avec le médecin. Il faut aussi créer un droit à la transparence, qui va avec la liberté de choix. Je pense qu'on doit être en mesure de s'opposer à être traité, analysé ou diagnostiqué par un robot, mais il faut savoir quelles conséquences cela implique. Enfin, on doit avoir une traçabilité de ce que fait le robot, afin de pouvoir comprendre, corriger si nécessaire, et peut-être sanctionner.

En ce qui concerne les données, il faut reconnaître, de manière radicale, que le monde de demain a besoin de l'intelligence artificielle. Dans le domaine de la santé, le progrès viendra de l'*open data*. Je crois que nous devons aller vers une ouverture généralisée des données de santé, sur la base d'un consentement libre et éclairé. Il y a des gens qui donnent leur sang, et ils ont bien raison de le faire car les besoins sont importants, mais il en existe peut-être aussi qui, comme votre serviteur, sont prêts à partager leurs données. Il faut une protection de la vie privée et de l'intimité, mais les données de santé sont nécessaires pour le développement des algorithmes, qui prendront alors de la puissance.

L'ouverture des données de santé correspond à un intérêt majeur : il faut déverrouiller non pas le secret médical – il n'en est pas question – mais les données liées aux maladies, grâce à un processus de pseudonymisation plutôt que d'anonymisation. Il y a un risque, que la société doit gérer au regard de l'intérêt des nouveaux systèmes.

Dernier élément en ce qui concerne l'*open data*, les données de santé se retrouvent partout. Il faut que l'individu puisse, à un moment donné, être le musée de ses propres données de santé, ce qui implique de pouvoir les récupérer.

Pour conclure, je voudrais souligner qu'une intelligence artificielle conçue pour être éthique *by design* est une nécessité dans le domaine de la santé. En l'inscrivant dans le droit, on arrivera peut-être à améliorer l'état de santé de tous, partout dans le monde. Je participe à un projet visant à mettre l'intelligence artificielle au service du monde entier : on peut changer le monde grâce à elle, et on va commencer par l'utiliser pour tous les enfants. Agissons maintenant. Nous avons déjà adopté des droits de l'homme fondamentaux, puis des droits de l'homme numériques, dans le cadre de la loi « informatique et libertés », qui est toujours pertinente, quarante ans plus tard ; créons maintenant les droits de l'homme de l'intelligence artificielle.

Mme la présidente Yaël Braun-Pivet. Une vaste tâche nous attend donc. (*Sourires.*)

M. Cédric Villani. C'était un plaisir d'écouter ces différentes contributions, qui m'ont replongé dans les nombreux débats auxquels j'ai participé dans le cadre du rapport au Gouvernement sur l'intelligence artificielle, que j'ai coordonné. Les trois interventions que nous venons d'entendre s'inscrivent parfaitement dans le cadre de ces débats mais aussi de notre analyse et de nos recommandations.

Je voudrais revenir, très rapidement, sur un certain nombre de points, en insistant tout d'abord sur le fait que l'intelligence artificielle n'a pas de définition précise, ce qui est bien embêtant si l'on veut travailler sur des dispositions législatives spéciales. Il s'agit de toute solution algorithmique servant à effectuer une tâche avec précision et efficacité – cela peut consister, par exemple, à reconnaître une pathologie ou à suggérer un diagnostic. L'intelligence artificielle est extrêmement efficace et elle s'invite dans tous les domaines. C'est un sujet multifacettes.

J'insiste toujours sur trois mots-clés lorsque je présente mon rapport. Le premier est l'expérience. Il est important de développer et d'améliorer les outils d'intelligence artificielle par l'expérimentation. C'est là-dessus qu'il y a actuellement les batailles les plus âpres au niveau international : chacun veut avoir les meilleurs chercheurs et les meilleurs appareils pour réaliser des expérimentations. Le deuxième mot-clé est la souveraineté : la question de l'intelligence artificielle est liée au contrôle de soi-même et de la société, ainsi qu'à la capacité à influencer sur l'avenir, dans un contexte de plus en plus concurrentiel. Le troisième mot-clé est le partage. Comme plusieurs intervenants l'ont très bien dit, c'est dans le partage des données, qu'il soit ouvert ou non, dans la façon de mettre ensemble les données et de chercher les corrélations que se trouve bien souvent la valeur ajoutée. C'est une solidarité nouvelle, à côté d'autres formes plus classiques telles que le partage de temps ou le partage du risque dans les démarches mutualistes.

Parmi les enjeux qui nous concernent directement, dans la perspective de la révision de la loi de bioéthique, il importe de ne pas se tromper d'adversaire quand il est question de l'intelligence artificielle. On pourrait limiter les expérimentations afin d'éviter des situations éthiquement difficiles mais, ce faisant, on empêcherait le développement de solutions efficaces, qui pourraient, dans certains cas, sauver des vies. Il importe donc de faire des expérimentations tout en les encadrant.

Nous avons besoin de théoriciens de l'algorithmique, mais aussi et surtout de données et de matériel. Il faut des données bien annotées et sûres, comme l'a dit M. Lucas, et l'on ne saurait trop insister sur ce point. Un autre sujet majeur est de savoir qui partage ses données avec qui, ce qui est souvent la question la plus difficile.

Chaque problème d'intelligence artificielle comporte trois défis : le défi scientifique et technologique, le défi éthique et juridique, et le défi de la confiance – c'est-à-dire « qui travaille avec qui ? », et « qui s'ouvre à qui ? » En pratique, presque toutes nos auditions ont montré que l'on surestime les défis scientifiques et technologiques, alors que l'on sous-estime le défi du partage. Il faut prendre en compte cette dimension dès lors que l'on s'intéresse à l'évolution de l'éthique ou du droit.

Le deuxième axe, après celui de la recherche, concerne l'efficacité. À quels problèmes cherche-t-on à appliquer l'intelligence artificielle ? À des problèmes qui nécessitent de prendre en compte beaucoup de paramètres et d'aboutir à une solution souple, évolutive et personnalisée. La santé est un cas d'école : elle dépend d'un nombre de paramètres considérable, les mécanismes sont souvent mal élucidés, la situation se personnalise en fonction de la pathologie et du médecin, et il y a des évolutions liées aux connaissances, à l'état de santé du médecin et à la culture. Le fait que l'on se focalise à chaque fois sur des tâches précises, dans un contexte donné, est un élément important pour ne pas se méprendre sur ce qu'est l'intelligence artificielle.

La question de l'ouverture des données a été évoquée, en particulier par Alain Bensoussan. Il y a un lien avec de nombreux sujets, comme celui de la confidentialité des données. Aux États-Unis, où des expériences ont pu être faites, une universitaire, Latanya Sweeney, s'est rendue célèbre en montrant qu'il était possible de réidentifier certains dossiers particuliers au sein de données anonymisées et ouvertes à des fins de recherche. Elle a publié un article expliquant comment on pouvait réidentifier le dossier du gouverneur du Massachusetts parmi des données prétendument anonymes. Ce genre de démonstration nous interroge sur ce qui est possible et ce qui doit être autorisé.

En matière d'éthique, la première question est bien sûr celle du respect de la vie privée. Nous avons récemment eu l'occasion de débattre du règlement général sur la protection des données (RGPD). En matière de responsabilité, l'intelligence artificielle se caractérise par le recours à des bases de données importantes, afin de guider l'action par l'exemple. Mais qui dit base de données, dit multiplicité d'ingrédients dans le produit, donc dilution des responsabilités, ce qui peut conduire à des difficultés pour établir la chaîne de responsabilité. Il y a aussi des difficultés en matière de propriété intellectuelle et de traçabilité des décisions.

De manière générale, on peut avoir du mal à anticiper les problèmes qui vont se poser. Une référence en la matière est l'ouvrage de l'universitaire américaine Cathy O'Neil, *Weapons of Math Destruction*, qui fait état de toutes sortes d'utilisations de l'intelligence artificielle posant problème du point de vue éthique – pas forcément sur le plan des grands principes, mais parce que l'utilisation, très efficace, de certains algorithmes peut conduire à des pratiques inacceptables, comme des manipulations – on a vu les débats qui ont eu lieu après certaines élections –, ou à des organisations humaines inacceptables, par exemple en termes de conditions de travail. Quand on regarde les usages de l'intelligence artificielle, il ne faut pas se limiter à des cas déterminés par avance : il faut laisser une marge d'évolution par rapport à la société et aux usages.

Mon rapport recommande de créer un comité d'éthique *ad hoc*, spécialisé, en parallèle de celui qui existe actuellement. Ce comité d'éthique se saisirait de toutes sortes de questions d'éthique et d'intelligence artificielle, afin de conseiller le Gouvernement et de répondre aux questions du Parlement et de la société en général. Il y aurait bien sûr une coordination avec l'actuel comité d'éthique, voire une possibilité de saisine conjointe. Une autre option, que nous n'avons pas recommandée mais qui est envisageable, serait que l'actuel comité d'éthique étende ses compétences au-delà des sujets de bioéthique. Il reviendra au professeur Delfraissy de formuler des recommandations dans le cadre de son rapport.

Une grande partie des enjeux concernant l'intelligence artificielle sont liés à la connaissance de notre propre humanité et à ce que nous voulons ou acceptons en tant que société, par rapport à notre identité. Ce sujet a été évoqué tout à l'heure lorsque M. Loiseau a insisté sur la question de l'homme « augmenté ». Très souvent, les questions que nous nous posons au sujet de l'intelligence artificielle nous renvoient en fait à nous-mêmes. Notre rapport avec certains assistants affectifs, par exemple, interroge beaucoup. Quand un assistant est explicitement conçu pour combler un vide affectif, pour prendre la place d'une compagne ou d'un compagnon absent, pour susciter une réponse que l'on aurait envie de faire à un humain, on peut se demander s'il n'y a pas une sorte de fraude à l'humanité. Alain Bensoussan a insisté sur l'importance de savoir si ce que l'on a en face de soi est un algorithme ou un humain. Même quand on sait qu'il s'agit d'un algorithme, on peut être tenté de le considérer comme un humain et de dévier peu à peu de notre humanité. À mesure que les usages se multiplieront, on pourra être amené à se demander ce que nous acceptons ou non, en tant que société, dans l'imitation de l'humanité par l'algorithmique, en ne perdant pas de vue que le but est de servir l'humain en rendant les pratiques beaucoup plus efficaces dans de nombreux domaines.

M. Thomas Mesnier. Je voudrais vous interroger sur le développement de l'intelligence artificielle en matière d'aide à la décision et au diagnostic médical. On peut aujourd'hui poser des diagnostics sur certaines pathologies grâce à des algorithmes et à des croisements extrêmement précis de données qui permettent de distinguer, par exemple, les tumeurs bénignes et les cancers. Même si la Haute Autorité de santé (HAS) prévoit déjà un encadrement des logiciels d'aide à la décision et à la prescription médicale, nous ne sommes qu'aux prémices du travail législatif dans ce domaine. Qui est responsable en cas de dysfonctionnement ou d'erreur de diagnostic ? De quelle manière l'expertise humaine doit-elle intervenir aux côtés de l'expertise artificielle ?

Mme Caroline Janvier. J'avais l'intention de vous interroger sur la responsabilité qui est engagée lorsqu'une opération ou une prise de décision fait intervenir un médecin et un système d'intelligence artificielle, mais Thomas Mesnier vient de le faire, d'une manière très complète.

M. Éric Diard. Merci pour ces interventions liminaires qui sont particulièrement intéressantes, mais aussi anxiogènes, du moins pour certains parlementaires, dont je fais partie. Les défis de l'avenir nous conduisent à nous interroger.

En ce qui concerne l'encadrement juridique, il est clair que le robot ne doit pas être un homme ou un animal. Pouvez-vous revenir sur cette question ? En revanche, l'intelligence artificielle peut être mise au service de l'homme en remplacement de certains animaux, par exemple pour rechercher des individus ou prévoir des événements tels que les tremblements de terre – des animaux en sont capables. Des cancérologues et des oncologues disent que des

animaux peuvent aussi détecter certains cancers grâce à leur odorat, et l'on envisage maintenant des nez artificiels qui permettraient de prévoir l'apparition de certains cancers. J'aimerais vous entendre également sur ce point.

Mme Charlotte Lecocq. Ma question porte plus particulièrement sur l'impact de l'utilisation de l'intelligence artificielle dans le monde du travail, notamment parce que cela implique une importante collecte de données. Qui va s'en charger ? Devra-t-on mobiliser spécifiquement des personnes pour accomplir cette tâche qui peut s'avérer assez pénible, si elle représente l'essentiel des missions exercées ? Cela implique-t-il, par ailleurs, une hyperconnectivité permanente des professionnels, sans véritable possibilité de se déconnecter ? Une collecte efficace des données renvoie un peu à la question du « flicage ».

M. Jacques Lucas. Rassurez-vous, monsieur Diard, vous n'êtes pas le seul à éprouver quelque angoisse : les citoyens ne veulent pas voir le lien entre la personne malade et le professionnel de santé perdre de son humanité. Mais ceux qui voudraient que rien ne change et qui refusent de voir le monde dans lequel nous vivons mettent toujours en avant cette déshumanisation. M'exprimant au nom de l'Ordre, il m'arrive de citer Hippocrate, qui a soustrait l'homme à la puissance des dieux en montrant que les sacrifices expiatoires, personnels ou collectifs, ne servaient pas à la guérison et que le médecin devait observer l'homme dans son biotope pour lui apporter les secours de la médecine. Il est indiscutable que le numérique, les technologies, l'intelligence artificielle, en évitant les discriminations, peuvent représenter un plus pour la bienveillance – une bienveillance républicaine.

Bien évidemment, des dangers existent, et le législateur doit poser les limites. La question qui inquiète l'ensemble des professionnels de santé est celle de leur responsabilité juridique. Jusqu'où sera-t-elle engagée, dès lors que l'intelligence artificielle parvient à donner des résultats de qualité supérieure ? C'est particulièrement vrai s'agissant de l'interprétation de clichés en imagerie médicale. L'ensemble de la profession radiologique a d'ailleurs décidé de s'engager dans la construction d'un outil d'intelligence artificielle qui collectera les clichés et les comptes rendus – les bases sont considérables, tant dans le monde hospitalier que dans le monde libéral. Ce genre d'outil n'existe pas sur le sol national.

L'autre question est de savoir jusqu'où les moyens d'intelligence artificielle seront déployés. Avec un outil doté d'une intelligence artificielle moyenne, un individu pourra saisir ses propres données et recevoir un diagnostic sans qu'aucun intermédiaire, médecin ou professionnel de santé, n'intervienne. Mais une personne peut-elle recevoir de façon brutale et sans aucune intermédiation humaine un diagnostic de mélanome, avec un pronostic de survie à dix-huit mois ? Les industriels, les fournisseurs, les sociétés peuvent-ils mettre directement sur le marché ces dispositifs ? Sans doute le législateur doit-il se prononcer sur cette question très importante.

C'est aussi le cas des dispositifs plus légers. Avec les objets connectés, l'individu se trouve bardé de capteurs qui enverront directement les données sur une base chargée de les collecter et de les traiter. Ces objets doivent-ils être considérés comme des dispositifs médicaux, avec les règles propres du marquage, ou doit-on envisager un processus de labellisation qui soutienne le marché industriel et garantisse aux utilisateurs la fiabilité technologique de la collecte et la protection de leurs données personnelles ? Voilà des questions anxiogènes.

M. Grégoire Loiseau. Il est vrai qu'en matière de responsabilité médicale, le droit actuel est dépassé. La loi de 2004 s'est bornée à entériner ce que la Cour de cassation avait

forgé au fil du temps. Les distinctions subtiles, s’agissant d’actes médicaux réalisés par des personnes humaines, n’ont plus cours depuis qu’un outil intelligent peut participer au diagnostic ou aux soins. Le législateur doit s’emparer du sujet sans trop attendre, car cela participe du phénomène d’acculturation sociale de l’intelligence artificielle.

Loin de vouloir alimenter l’anxiété, je pense que l’acculturation sociale passe par des solutions sûres. Cela n’implique pas de revisiter tout le droit – je ne crois pas qu’il soit nécessaire de refondre entièrement notre système juridique –, mais de repenser certains volets. La responsabilité médicale doit être adaptée aux nouveaux outils, en particulier aux outils intelligents qui interviennent soit dans le diagnostic, soit dans l’acte médical lui-même, quitte – pourquoi pas ? – à dissocier la responsabilité du médecin, être humain, de celle de la machine, selon que l’acte médical a été accompli par l’un ou par l’autre.

Lorsque le médecin humain aura posé un diagnostic inexact, on aura tendance à le lui reprocher, considérant qu’il s’est trompé là où la machine avait raison. Autrement dit, la plus grande fiabilité du diagnostic posé par l’outil intelligent pourrait conduire à surinvestir ou à augmenter le risque de responsabilité portant sur les médecins. C’est la raison pour laquelle j’évoque l’éventualité d’un régime dissocié, adapté aux nouvelles fonctions de l’intelligence artificielle en matière médicale. Il est certain que, sur ce point, notre droit doit évoluer.

Il est tout aussi clair que la loi Badinter de 1985 n’est pas adaptée à la mise en circulation des voitures autonomes. Lorsque l’on aura dépassé le stade de l’expérimentation et que des véhicules autonomes seront mis en circulation sur les voies ouvertes au public, il faudra adapter cette loi conçue il y a plus de trente ans pour des véhicules conduits par des êtres humains. Le législateur doit procéder à une adaptation du droit, mais progressivement, sans tout révolutionner.

Au risque de paraître passéiste ou conservateur – ce que j’assumerais volontiers –, je suis avant tout un humaniste, qui tient l’être humain pour valeur fondamentale. C’est au travers du grand principe posé par vos prédécesseurs en 1994 – la loi garantit la primauté de la personne –, que j’aborde les questions du statut juridique des animaux et du statut juridique de l’intelligence artificielle. N’en déplaise à Alain Bensoussan, « personne » n’est pas un mot-valise, « personne » n’est pas une enveloppe abstraite et juridique. Non, elle recouvre des valeurs. La distinction entre les personnes et les choses remonte à la philosophie grecque et continue de faire sens aujourd’hui.

Le législateur a imité la notion en 2015 en écrivant que l’animal est un être doué de sensibilité et que, sous réserve des lois qui le protègent, il est soumis au régime des biens. Il a eu l’intelligence d’insérer cet article non pas à la fin du livre I du code civil traitant des personnes, mais au tout début du livre II, intitulé « *Des biens et des différentes modifications de la propriété* ». L’animal est une chose, mais une chose protégée, qu’il convient de respecter, sans qu’il soit nécessaire d’en faire une « personne ».

Pourquoi ne pas raisonner de la même manière pour certains objets intelligents ? Des relations émotionnelles – et parfois bien plus que cela –, peuvent se nouer entre eux et les êtres humains. Il faut rappeler qu’il existe une hiérarchie entre les êtres humains, les autres êtres et les choses. J’assume pleinement ce propos que je sais ne pas être forcément majoritaire, mais cette hiérarchie fonde notre droit depuis l’Antiquité, et bien que cela fasse preuve de passéisme, j’y suis attaché.

M. Alain Bensoussan. Je voudrais dépasser le débat sur la « personne » robot. Lorsque je l'ai lancé au niveau mondial, il est clair que le terme « personne » ne renvoyait pas à l'humain ; les robots ne sont pas des humains-moins, ce ne sont pas des objets-plus, ce ne sont pas davantage des animaux, encore moins des enfants. Oublions donc le mot ; il s'imposera par la pratique, parce qu'il est simple.

Regardons ce que cela signifie exactement en matière de responsabilité. Une société qui entre dans un choc d'intelligence artificielle est une société qui va déterminer un régime démocratique de la responsabilité artificielle.

Aujourd'hui, le taux d'erreur d'une voiture autonome est de 10^{-9} , le même qu'en matière d'aviation. Chaque civilisation déterminera le niveau de risque acceptable. Il s'agit d'un autre schéma : faire au mieux pour l'intérêt de l'humanité, avec un minimum de risques, mais en se déplaçant vers cette notion de mieux. Dans le monde de demain, c'est ce concept qui différenciera les sociétés ou les nations.

Les lois de demain renverront à des référentiels, et ces référentiels seront d'autant plus serrés que l'on placera la vie humaine au-dessus d'autres avantages. Il semblerait que l'accident de la voiture Uber soit dû au fait que le véhicule a été réglé de telle manière que l'on prenait un risque : on freinait moins pour maximiser le confort. J'use de beaucoup de précautions, mais il semble bien qu'il s'agisse d'un problème de réglage de ce « thermostat de la responsabilité ».

Il faut donc placer le niveau de responsabilité technologique au plus haut. L'intérêt est que cela devient calculable. On peut le répéter, l'enfermer dans une enveloppe de potentialités. N'oublions pas que nous sommes en présence d'un cerveau artificiel primitif, parce qu'il y a apprentissage, expérience et mutation.

Le schéma de responsabilités que je voudrais vous proposer fixe la règle démocratique du taux d'erreur acceptable. Il faut sortir l'intelligence artificielle, les algorithmes, les *data* des dispositifs médicaux, mais réintroduire tout de suite, comme vous l'avez dit, un système de certification, de labellisation, de telle manière que l'on augmente le taux de confiance. Les responsabilités sont en cascade : responsabilité de la plateforme d'intelligence artificielle – c'est-à-dire des algorithmes – responsabilité des données, responsabilité de l'utilisateur. J'appelle votre attention sur le fait que l'utilisateur a un rôle à jouer, dans sa capacité à permettre cet apprentissage.

Pour développer la confiance, il faut non pas confronter le robot à l'humain, mais confronter le robot à la qualité de ses diagnostics, notamment médicaux. C'est quelque chose de mesurable. J'appelle de mes vœux la création d'une Agence de l'intelligence artificielle – un tribunal de l'intelligence artificielle – comme cela existe en matière d'énergie nucléaire, avec une procédure spécialisée. En construisant, au moins de manière expérimentale, une responsabilité juridique de l'intelligence artificielle, on créera la confiance et l'acceptabilité nécessaires, car le monde de demain sera un monde guidé par l'intelligence artificielle.

Mme la présidente Yaël Braun-Pivet. Vous avez dit que la loi doit garantir la primauté de la personne. J'entends qu'il faut placer la vie humaine au-dessus de tout. Les juristes que vous êtes considèrent-ils que cela devrait faire l'objet d'une incise constitutionnelle, en ces temps de révision ?

Mme Typhanie Degois. Pascal Picq, dans son dernier livre paru en 2017 *Qui va prendre le pouvoir ? Les grands singes, les hommes politiques ou les robots ?*, explique pourquoi l'Europe est en retard sur les robots, notamment par rapport à l'Asie. Notre société demeure méfiante et hésite à laisser une place, à donner des droits, que ce soit aux animaux ou aux robots, parce qu'elle veut conserver aux humains leur supériorité.

En matière d'intelligence artificielle, les derniers progrès sont dans l'apprentissage automatique, le *machine learning*. Désormais, on avance même avec des réseaux de neurones artificiels, des *deep learning*. Les technologies deviennent proactives, elles apprennent par elles-mêmes.

Parallèlement, il existe une opacité autour des algorithmes. On peut se demander aujourd'hui s'ils sont éthiques. Peut-on anticiper la logique d'une *machine learning*, comme ce qui s'est passé avec la voiture Uber ? Peut-on garder un contrôle sur ces technologies ? Selon vous, dans quelle mesure peut-on encadrer les algorithmes ?

Mme Mireille Robert. En cas d'erreur du robot, qui sera tenu responsable ? Faut-il créer un statut juridique propre aux robots ? Devons-nous les doter, à l'instar des entreprises ou des associations, d'une personnalité juridique de type particulier ? Ou voulons-nous continuer à les considérer comme des objets appartenant à une personne responsable de leur utilisation ?

M. Jean-Pierre Pont. Vous avez parlé d'un droit fondamental, le droit à la connaissance. Ne risque-t-on pas de voir des personnes qui ne comprennent pas ce droit ou qui utilisent mal cette connaissance ? Désormais, dans les cabinets médicaux, des patients se présentent déjà munis d'un diagnostic, sur lequel il devient très difficile pour le médecin de revenir.

Mme Maina Sage. Ce débat très riche nous rappelle tout d'abord à notre propre responsabilité d'élus, qui devront légiférer sur ces sujets. Ce sont des questions qui nous interpellent personnellement, impliquent une définition de l'humanité, alors que nous ne mesurons pas encore tous les effets du développement des nouvelles technologies sur la société. Nous observons une forme de déshumanisation de la société, de la jeunesse notamment.

Quel est le risque que la *e-santé* et l'utilisation de l'intelligence artificielle dans le champ de la santé soient un cheval de Troie pour faire accepter à l'humanité ces nouvelles techniques ? Pour la Polynésie française, territoire lointain et très étendu, la télémédecine est un enjeu majeur, une chance de développement et de sécurité pour les habitants. Nous suivons donc de près ces évolutions. Mais jusqu'où le mieux est-il l'ami du bien ? Jusqu'où pourra-t-on préserver le sens de l'humanité ?

Le développement des technologies ne doit pas se faire au détriment d'autres parties du monde. Est-il encore possible d'ailleurs de contenir cette évolution ? Le sujet est nouveau en Europe, mais il est déjà bien développé dans d'autres régions du monde, où les notions éthiques ne sont pas tout à fait les mêmes. Je crains que l'Europe ne se projette dans l'intelligence artificielle pour rester compétitive sur le plan mondial, et oublie l'éthique.

Mme Blandine Brocard. Je n'ai aucune connaissance en la matière, et je rejoins nos collègues qui trouvent tout cela anxiogène. Je comprends que le système est « apprenant », « autonome » et « mutant » – des termes qui, lorsque l'on n'est pas de la partie, font un peu

peur. Toutefois, vos interventions permettent d'apaiser quelque peu ce sentiment et de nous éclairer. Il faut que nous gardions à l'esprit que l'objet de l'intelligence artificielle est d'être à notre service. Mais comment s'assurer qu'elle le reste ? Peut-on être certain que nous ne jouons pas aux apprentis sorciers ?

Je me demande aussi dans quelle mesure l'intelligence artificielle, qui risque de profiter à certains seulement, ne contribuera pas à fracturer encore davantage notre société, voire notre pays.

Cédric Villani a expliqué tout à l'heure, et il faut le répéter, le marteler, que l'on ne peut imiter notre humanité. Monsieur Loiseau, je trouve du réconfort dans vos propos, puisque vous dites que l'être humain est la valeur fondamentale et qu'il faut réaffirmer la hiérarchie qui fonde notre droit. Je partage votre vision, qui est loin d'être passéiste ! Mais comment s'assurer que l'être humain demeure la valeur fondamentale ? Hélas, l'homme finit par être dépossédé de ses inventions, même les plus louables.

M. Christophe Lejeune. Google vient d'annoncer à ses employés qu'il abandonnait son projet de partenariat avec le ministère de la défense américain. Ses détracteurs expliquent que ce projet de recherche en intelligence artificielle devait permettre de repérer, grâce à des drones, les corps inertes sur les champs de bataille. Cela va à l'encontre de ce que l'on peut espérer d'une intelligence artificielle au service de l'homme. Pensez-vous que Google a cessé le partenariat dans le domaine militaire par peur de ce qu'il pouvait en ressortir à terme, ou pensez-vous que, de manière beaucoup plus mercantile, la société a cédé aux lobbies qui préféreraient la voir privilégier les intérêts commerciaux à dominante civile ?

M. Brahim Hammouche. Les glissements sémantiques concernant la personne humaine me gênent. Je pense qu'il faut réserver à l'humain les notions d'intentionnalité, car il n'existe pas d'intentionnalité robotique, de liberté, d'erreur – il y a une vertu à en commettre – et de responsabilité. Il convient aussi de rappeler que l'intelligence, dans notre société, est plurielle : je pense à l'intelligence du cœur, à l'intelligence collective, à l'intelligence de la main, prolongement de l'esprit.

L'idéal d'une vie sans erreurs, sans accidents, n'est-il pas le produit d'une vision perfectionniste, donc inaccessible ? Ne risque-t-on pas de s'éloigner de l'essentiel, qui fonde notre rapport aux êtres, aux choses, au monde – je veux parler de la relation ?

Mme Delphine Bagarry. Non seulement je suis angoissée, mais j'ai le vertige ! J'ai bien saisi l'enjeu de partage et de solidarité à l'échelle mondiale que soulève l'intelligence artificielle.

En vous entendant expliquer que le robot, en médecine, ne commet d'erreur qu'une fois sur dix, lorsque l'homme se trompe une fois sur deux, je me suis demandé si, dans le domaine de la loi, il ne serait pas également supérieur à l'homme. Ne serait-il pas plus à même que nous de rédiger ces lois de bioéthique ? C'est une question vertigineuse.

Mme la présidente Yaël Braun-Pivet. Messieurs, vous avez la parole pour répondre à cette dernière série de questions.

M. Jacques Lucas. L'intelligence artificielle permettra à des régions ou à des collectivités, mais également à tous les pays émergents, de résoudre les problèmes sanitaires majeurs auxquels ils sont confrontés. Il m'est toujours resté en mémoire le mot d'un confrère africain lors d'un colloque sur le consentement : « *Pour consentir aux soins, encore faut-il ne*

pas être mort faute de les avoir reçus ! ». Et sans vouloir me prendre pour le Saint-Père à l'ONU, je vous dirai : « *N'ayez pas peur !* » Il faut analyser les situations et combattre les risques, mais se garder de mettre en avant les risques par rapport aux avantages conséquents que l'intelligence artificielle peut apporter, notamment au reste du monde.

Certes, le législateur français ne peut prétendre faire la loi du monde, mais, porteur des vieilles valeurs républicaines, il peut être moteur dans la prise de conscience collective européenne. L'intérêt ne peut être simplement celui du marché, nous devons porter des valeurs éthiques.

Je me permettrai de m'abriter derrière un vice-président du Conseil d'État, M. Renaud Denoix de Saint Marc, qui estimait en 2001 que la loi devait être « *solennelle, brève et permanente* », au lieu d'être « *bavarde, précaire et banalisée* ». La question des technologies n'est pas directement inscrite dans les aspects de bioéthique mais elle pourrait être incluse dans la révision des lois de bioéthique. Comme l'a dit le président du Comité consultatif national d'éthique, compte tenu de la rapidité des évolutions technologiques, il est très difficile de tout anticiper.

La loi doit être brève, solennelle et permanente ; elle doit être fondée sur les grands principes de non-discrimination et d'irréductibilité de la personne. Le malade est d'abord une personne, ce n'est pas un individu auquel on va soustraire ses *datas* pour donner un diagnostic. Car après tout, la conduite à tenir vis-à-vis de la maladie diagnostiquée dépendra de ce que souhaite le malade : on pourra lui proposer diverses alternatives thérapeutiques. Ainsi, une femme pourra dire : « *Je ne veux pas me faire opérer du cancer du sein, bien que le docteur Watson ait dit que c'était ce qu'il fallait faire* ». Le médecin pourra alors essayer de la convaincre et, en cas d'échec, lui proposer une alternative thérapeutique.

Si la loi est brève et solennelle, il faut qu'existe un organisme de régulation. Nous sommes tout à fait favorables à ce que l'on produise du droit souple, à l'image de la *soft law* anglo-saxonne. Le droit français n'y est pas très favorable ; le Conseil d'État lui-même est divisé entre les tenants du droit dur et les auteurs du rapport de la section du rapport et des études, qui avaient théorisé le droit souple. Si vous légiférez, il conviendrait que la loi comporte la possibilité de réguler par du droit souple.

Reste à savoir qui régulera. Cédric Villani a évoqué tout à l'heure un comité *ad hoc* « parallèle » au Comité consultatif national d'éthique. Il pourrait être plus opportun de prévoir que le CCNE comporte une chambre qui siégerait sur les questions de santé et sciences de la vie et une chambre qui travaillerait sur les questions relatives au numérique. Ces deux chambres pourraient se réunir sur les sujets de santé et d'intelligence artificielle.

Je ne pense pas que l'intelligence artificielle remplacera un jour le législateur, je ne crois pas non plus que nous éliminerons des robots. En revanche, les métiers juridiques seront fortement impactés. Une machine qui aura intégré toute la jurisprudence pourra un jour dire à un justiciable les chances qu'il a de gagner une affaire, sans qu'il soit besoin de plaider. Je sais que la Cour de cassation, notamment, s'intéresse beaucoup à ce sujet.

M. Grégoire Loiseau. Je ne répéterai pas ce qui a été très bien dit. Faut-il légiférer et sur quoi ? Il faut être extrêmement prudent. Nous sommes au début d'un phénomène dont on perçoit parfois plus les risques que l'on en voit les avantages, et la situation est en effet anxiogène. Je crois qu'il faut se laisser du temps pour appréhender le phénomène et ne pas surlégiférer, au risque de figer les choses et de produire un droit autarcique.

L'époque où la France exportait son code civil est révolue. Aujourd'hui, si l'on veut faire un droit qui ait du sens, il doit être régional, au moins européen. Compte tenu des enjeux techniques, économiques et de droit social, rien ne se fera d'utile, d'efficace et d'efficient qui serait purement national. Cela ne signifie pas qu'il faille rester les bras ballants, en attendant que la Commission européenne et le Conseil avancent des projets. La France a un rôle moteur, parce que son influence au sein de l'Union est historique et aujourd'hui très importante, mais aussi parce qu'elle s'apprête à légiférer sur certains points.

La plume ne doit pas être bavarde, mais simple. Il s'agit de montrer l'exemple en rappelant quelques grands principes. Madame la présidente, je ne suis pas convaincu qu'inscrire dans la Constitution le principe de la primauté de l'être humain ou, ce qui reviendrait au même, celui de dignité humaine fera avancer juridiquement les choses, mais poser ce principe de hiérarchie dans le texte fondamental de notre République serait un acte symbolique fort.

Donnons-nous le temps. Comme cela vient d'être dit, le droit mou, auquel nous ne sommes pas accoutumés, peut être extrêmement utile dans cette période de latence, qui est aussi période d'appropriation, d'acculturation du phénomène. Aujourd'hui, on n'ose plus parler de morale et l'on préfère aller chercher chez les Grecs l'éthique, sans doute plus rassurante. Au moins à titre transitoire, l'éthique peut avoir eu une utilité. Alors que les acteurs eux-mêmes, comme Google ou Facebook, travaillent sur des chartes internationales éthiques, servant évidemment leurs intérêts, il faut que les instances nationales, supranationales, dans toute leur légitimité, posent cette éthique. Nous avons la responsabilité de le faire, sans attendre la charte éthique Facebook, déjà très avancée, au risque de nous la voir imposer, sans d'autre choix, d'un point de vue commercial et de compétitivité, que de s'y soumettre.

Il convient donc de légiférer avec beaucoup de prudence, plutôt dans le sens de l'acculturation sociale, c'est-à-dire uniquement quand cela est nécessaire, et d'utiliser l'éthique à titre transitoire. Il faut aussi agir au niveau de l'Union européenne pour faire valoir les grandes valeurs d'humanité qui ont fait les démocraties européennes, rappeler que la technologie est au service des humains et démonter le mythe du robot maître du monde. Nous en sommes très loin pour le moment, avec une intelligence artificielle de type faible ou moyen. L'idée d'une prise de pouvoir par les robots est, pour le coup, inutilement angoissante.

M. Alain Bensoussan. Il s'agit d'abord d'une problématique mondiale. Les robots sortent des laboratoires ; ils sont dans les hôpitaux, dans les usines, dans les entreprises, dans les domiciles. Cela signifie que nous sommes aujourd'hui dans un régime de liberté. Toute personne qui a envie de faire de l'algorithmie d'intelligence artificielle peut le faire en toute liberté, c'est-à-dire sans conscience, sans contrôle, sans responsabilité, sans assureur et sans dignité. Voilà la situation exacte. Nous sommes poussés, et si nous n'intervenons pas, nous trouverons demain face à un marécage.

Soit nous considérons que nous sommes dans un régime de liberté, que tout est permis, et que nous verrons ce que l'avenir nous réserve ; soit nous décidons de légiférer, de manière expérimentale, à petits pas.

Sur quoi faut-il légiférer ? J'aurais tendance à dire qu'inscrire dans la Constitution la primauté de la personne humaine est nécessaire et suffisant. Juridiquement, cela n'apportera pas grand-chose – encore que... –, mais pédagogiquement, en termes d'acceptabilité sociale, cela évitera le conflit des espèces, le conflit de la personne. Car lorsque l'on discute de la

personne-robot, on ne discute pas des avantages du robot sur la personne. Oublions les mots, dépassons ce clivage ! Le droit n'est que l'écume des valeurs, ce sont les valeurs qui importent.

La loi « Informatique et libertés », qui est devenue un standard mondial – alors que l'on pensait qu'il était impossible d'instaurer un standard mondial dans ce domaine – commence par une déclaration de principe : l'informatique doit être au service de chaque citoyen. De la même manière, il s'agirait d'écrire que l'intelligence artificielle doit être au service de l'humain, dans une approche humaniste et universelle. Voilà la première règle. Et si, demain, elle se trouvait dans la Constitution, ce serait merveilleux, car ce sont d'abord des valeurs que le droit traduit.

Au-delà, comment peut-on encadrer les algorithmes ? Il s'agit clairement et avant tout de poser le principe de dignité humaine, pour éviter des algorithmes malfaisants, pour éviter de tromper les gens, pour dire l'état des connaissances. Lorsque vous êtes opéré, vous recevez un document qui vous présente les risques. Il faudra y inclure un paragraphe sur l'intelligence artificielle, sur l'outil qui aura une emprise sur votre corps, sur le positionnement du médecin par rapport à l'intelligence artificielle dans l'acte médical.

Il s'agit ensuite d'une question de transparence : il est nécessaire d'informer sur les risques. Il faut aussi dire que l'on ne trompe pas, c'est le principe de loyauté. Je vous propose, dans ce cadre, de souligner un droit pénal spécial : il s'agit de dire que l'intelligence artificielle ne doit pas trahir les humains. Enfin, il faut prévoir la certification.

Vous avez évoqué la fracture intelligente. Non, le monde de demain ne sera pas un monde de robots remplaçant des humains, mais un monde d'humains augmentés remplaçant des humains. Le risque que la fracture de l'intelligence artificielle accentue la différenciation entre les humains est réel. Le problème n'est pas entre robots et humains, il est entre humains et humains.

Je veux aussi vous dire que le monde entier est en train de légiférer. Je pensais avoir inventé le système de cascade de responsabilités : pas du tout, une loi du Nevada a déjà posé la responsabilité de la plateforme d'intelligence artificielle et celle des fabricants ! À ce sujet, savez-vous pourquoi l'on parle de voiture autonome ? Tout simplement parce que le législateur avait inscrit dans la loi les mots *artificial intelligence platform*, avant de considérer que c'était trop dangereux en termes d'acceptabilité sociale. Il a alors choisi de remplacer les termes par *autonomous vehicle*. Vous voyez que derrière les mots, il y a du droit, et derrière le droit, des valeurs.

Exportons nos valeurs par une règle de droit qui se veuille humble. Cela permettra de baisser le niveau d'anxiété, d'augmenter le niveau pédagogique et de libérer l'intelligence artificielle dans la santé pour le bien du monde. Ce sera un grand pas pour tous.

Mme la présidente Yaël Braun-Pivet. Messieurs, je vous remercie infiniment de votre venue : je suis impatiente d'assister aux prochaines auditions, tant ces sujets sont passionnants.

La séance est levée à douze heures trente.

Présences en réunion

Réunion du mercredi 6 juin 2018 à 11 heures

Présents. – Mme Delphine Bagarry, M. Belkhir Belhaddad, M. Bruno Bilde, M. Julien Borowczyk, Mme Brigitte Bourguignon, Mme Blandine Brocard, M. Sébastien Chenu, M. Gérard Cherpion, M. Guillaume Chiche, Mme Josiane Corneloup, M. Dominique Da Silva, M. Pierre Dharréville, M. Jean-Pierre Door, Mme Audrey Dufeu Schubert, Mme Nathalie Elimas, Mme Catherine Fabre, Mme Caroline Fiat, Mme Agnès Firmin Le Bodo, Mme Emmanuelle Fontaine-Domeizel, Mme Albane Gaillot, Mme Patricia Gallerneau, Mme Carole Grandjean, M. Jean-Charles Grelier, M. Brahim Hammouche, Mme Monique Iborra, M. Cyrille Isaac-Sibille, Mme Caroline Janvier, Mme Fadila Khattabi, M. Mustapha Laabid, Mme Charlotte Lecocq, Mme Geneviève Levy, M. Gilles Lurton, M. Sylvain Maillard, M. Thomas Mesnier, M. Bernard Perrut, Mme Michèle Peyron, M. Laurent Pietraszewski, M. Adrien Quatennens, M. Alain Ramadier, M. Jean-Hugues Ratenon, Mme Mireille Robert, M. Adrien Taquet, M. Jean-Louis Touraine, Mme Isabelle Valentin, M. Boris Vallaud, Mme Michèle de Vaucouleurs, M. Olivier Véran, M. Francis Vercamer, Mme Annie Vidal, Mme Corinne Vignon, Mme Martine Wonner

Excusés. - Mme Ericka Bareigts, Mme Gisèle Biémouret, M. Paul Christophe, Mme Jeanine Dubié, Mme Claire Guion-Firmin, M. Thierry Michels, Mme Nadia Ramassamy, Mme Nicole Sanquer, Mme Élisabeth Toutut-Picard, Mme Hélène Vainqueur-Christophe, M. Stéphane Viry

Assistait également à la réunion. - Mme Caroline Abadie, Mme Laetitia Avia, M. Erwan Balanant, M. Ugo Bernalicis, M. Philippe Berta, M. Florent Boudié, Mme Yaël Braun-Pivet, M. Éric Ciotti, M. Jean-Michel Clément, M. Gilbert Collard, Mme Typhanie Degois, M. Vincent Descoeur, M. Éric Diard, Mme Nicole Dubré-Chirat, M. Jean-François Eliaou, M. Christophe Euzet, Mme Élise Fajgeles, Mme Isabelle Florennes, M. Raphaël Gauvain, Mme Marie Guévenoux, M. David Habib, M. Sacha Houlié, Mme Catherine Kamowski, M. Guillaume Larrivé, M. Philippe Latombe, Mme Marie-France Lorho, M. Olivier Marleix, M. Jean-Louis Masson, M. Fabien Matras, M. Stéphane Mazars, M. Jean-Michel Mis, M. Paul Molac, M. Pierre Morel-À-L'Huissier, M. Jean-Philippe Nilor, Mme Danièle Obono, M. Didier Paris, M. Jean-Pierre Pont, M. Éric Poulliat, M. Bruno Questel, M. Rémy Rebeyrotte, M. Robin Reda, M. Thomas Rudigoz, Mme Maina Sage, M. Hervé Saulignac, M. Jean Terrier, Mme Cécile Untermaier, M. Manuel Valls, M. Arnaud Viala, Mme Laurence Vichnievsky, M. Cédric Villani, M. Jean-Luc Warsmann, Mme Hélène Zannier